

Nina Gierasimczuk

Psychologiczna adekwatność modelu identyfikowalności w granicy

W jaki sposób człowiek uczy się swego pierwszego języka? Jak to możliwe, że korzystając ze szczątkowych i często wadliwych danych udaje mu się wykryć strukturę języka na tyle poprawnie, by móc się skutecznie porozumiewać? W jaki sposób zdobywa umiejętność konstruowania zdań, z którymi nigdy wcześniej się nie zetknął? Jak wreszcie opisać fragment umysłu odpowiadający za zdolności językowe, aby wyjaśnić fakt, że człowiek równie dobrze może jako pierwszego języka nauczyć się chińskiego jak i polskiego? Z tymi pytaniami mierzył się w swoich badaniach Noam Chomsky — twórca współczesnego paradygmatu badań nad językiem. Postulował on istnienie wrodzonego mechanizmu umysłowego, umożliwiającego nabywanie zdolności językowych. W pracach *Syntactic Structures* oraz *The Logical Structure of Linguistic Theory* zaproponował generatywistyczny model języka, który umożliwia analizowanie składni.

1. Lingwistyka Chomsky’ego

Istotne z naszego punktu widzenia cechy lingwistyki Chomsky’ego można określić następująco: rozróżnienie na semantykę i składnię w badaniach nad językiem oraz podkreślenie związków językoznawstwa z psychologią i filozofią umysłu.

Rozróżnienie na semantykę i składnię w badaniach nad językiem

Rozróżnienie na semantykę i składnię wprowadzone zostało do językoznawstwa przez Chomsky’ego. We wczesnym stadium rozwoju językoznawstwa semantyka i syntaktyka albo nie były rozróżniane albo granica pomiędzy nimi nie była wyznaczona (Lyons 1977).

Noam Chomsky właściwie nie zajmował się zagadnieniami semantyki języka naturalnego. Ścisłe zdefiniowanie pojęć składni pozwoliło jednak na precyzyjne wyznaczenie obszaru syntaktyki i w ten sposób oddzielenie go od kwestii znaczeń językowych. Wprowadzenie ścisłego, matematycznego

aparatu służącego badaniom składni stworzyło możliwość ujmowania języka naturalnego jako zbioru (przynajmniej potencjalnie) nieskończonego. Pozostając w zgodzie z poprawnością gramatyczną można bowiem dowolnie wydłużać zdania języka naturalnego, np.:

- (1) Genji popatrzył na swoich ludzi.
- (2) Piękny Genji popatrzył na swoich ludzi.
- (3) (...), posępny, piękny Genji popatrzył na swoich ludzi.

albo:

- (4) Posępny, piękny Genji popatrzył na swoich ludzi, którzy spali (...).

W ten sposób uzyskamy nieskończony zbiór poprawnie zbudowanych zdań. Istnieją oczywiście praktyczne ograniczenia na długość zdań. W ogólnym przypadku nie wiadomo jednak, w którym miejscu granicę tę umieścić. Dlatego, do opisu języka potrzebujemy narzędzia, które (potencjalnie) będzie produkować nieskończony zbiór zdań. Takim narzędziem jest gramatyka, która określa dopuszczalne sposoby łączenia elementów słownika. Dla potrzeb modelu zakładamy, że słownik jest ustalony, skończony i niezmienny. Model ten nie wyjaśnia więc zjawisk związanych z ewolucją języków — w tym paradygmacie zawsze badamy język w pewnym jego stadium. Przyjmujemy ponadto, że do opisu języka wystarcza skończenie wiele podstawowych reguł gramatycznych o charakterze rekurencyjnym. Wszak zdania 1–3 powstają wskutek powtarzania wciąż tej samej operacji — rozszerzania frazy rzeczownikowej przez dodanie przymiotnika z lewej strony.

Podkreślenie związków z psychologią, epistemologią i filozofią umysłu Zainicjowany przez Chomsky’ego przełom można wyrazić w stwierdzeniu: lingwistyka jest częścią psychologii, bowiem badania nad językiem przyczyniają się do głębszego zrozumienia procesów umysłowych (zob. np. (Chomsky 1968)). Moduł umysłu, który umożliwia posługiwanie się językiem cechuje się dużą dostępnością badawczą. Język jako zewnętrzny wytwór procesów umysłowych staje się ważnym przedmiotem dociekań psychologicznych (zob. (Pinker 2000), (Macnamara 1986)).

W swej teorii języka Noam Chomsky wyraża przekonanie, że jesteśmy wyposażeni w szereg swoistych zdolności (którym nadajemy nazwę „umysłu”), odgrywających istotną rolę w przyswajaniu wiedzy i umożliwiających nam działania dowolne, niezeterminowane przez zewnętrzne bodźce (choć podlegające pewnym wpływom z ich strony) (Chomsky 1965). Jedną z tych zdolności ma być możliwość przyswojenia języka, aktywowania tzw. kompetencji językowej. Warunkiem uaktywnienia tej wrodzonej potencjalności jest

obecność odpowiedniego bodźca. W jednej ze swych prac cytuje Leibniza: „idee i prawdy są nam wrodzone jako skłonności, dyspozycje, nawyki i naturalne potencjalności”, od siebie dodaje: „doświadczenie służy wydobywaniu, a nie formowaniu tych wrodzonych struktur” (Chomsky 1959). Jednym z celów lingwistyki Chomsky’ego jest odkrycie struktury ludzkiej kompetencji językowej wraz z jej wrodzonymi elementami. Teza o wrodzoności pewnych zdolności językowych wspierana jest następującymi argumentami.

Po pierwsze, językiem posługujemy się w sposób twórczy. Ze względu na nieskończony zbiór zdań, który może być generowany przez człowieka, języka nie można traktować ściśle behawiorystycznie, jako „nawyku” lub „zbioru dyspozycji do reagowania”. Zbiór zdań używanych przez człowieka nie zawiera się w zbiorze zdań wcześniej przez niego zasłyszanych. Jak już wcześniej wspomnieliśmy, postulowana przez Chomsky’ego kompetencja językowa ma postać zbioru reguł rekurencyjnych. Podmiot, wyposażony w zestaw takich reguł jest w stanie stwierdzać poprawność wyrażen lub tworzyć nowe wyrażenia, bez względu na to, czy zetknął się z nimi wcześniej czy nie. Ten aspekt lingwistyki Chomsky’ego sprawił, że określa się ją mianem generatywnej.

Po drugie, według Chomsky’ego, struktura języka ma charakter uniwersalny. Chomsky twierdzi, że tzw. „struktury głębokie” są bardzo podobne w różnych językach, a reguły rządzące przekształceniami „struktur głębokich” w „struktury powierzchniowe” dają się zamknąć w wąskiej klasie operacji formalnych (Chomsky 1995). Zjawisko to początkowo wyjaśniane jest przez odwołanie się do wrodzonego mechanizmu nabywania kompetencji językowej (Chomsky 1965). W późniejszych pracach Chomsky zakłada istnienie podstawowego zrębu wszystkich gramatyk, obecnego w każdym umyśle. Ten podstawowy zbiór reguł nazywa gramatyką uniwersalną (Chomsky 1995). Chomsky odrzuca inne możliwe przyczyny strukturalnego podobieństwa języków: postulat maksymalnej prostoty, do której mogły dążyć języki podczas swojego rozwoju, postulat „powszechnej logiczności” oraz koncepcję wspólnego pochodzenia wszystkich języków ludzkich. Nie tylko fakt podobieństwa różnych istniejących języków naturalnych jest dla niego zastanawiający. Równie zaskakujące jest podobieństwo występujące wśród zdolności językowych poszczególnych użytkowników tego samego języka.

[...] dany schemat gramatyki określa klasę możliwych hipotez; metoda interpretacji umożliwia testowanie każdej hipotezy na danych wejścia; miara ewaluacyjna wybiera najwyżej ocenianą gramatykę zgodną z danymi. Od momentu, gdy hipoteza [...] zostanie wybrana, uczący się zna język przez nią definiowany. W szczególności [...] [jest

on w stanie generować] nieskończony zbiór zdań z którymi nigdy się nie zetknął. (Chomsky 1967)

Wrodzonym elementem kompetencji językowej ma więc być według Chomsky’ego aparat służący nabywaniu zdolności językowej, *LAD* (skrót od *Language Acquisition Device*)(Chomsky 1965). Problem opisu *LAD* podjął E. M. Gold w pracy *Identification in the Limit*. Zaproponował tam model nabywania wiedzy składniowej, zwany identyfikacją w granicy. Postulowane przez niego rozwiązanie ma charakter hipotetyczny, jest jedną z możliwych odpowiedzi na problem postawiony przez Chomsky’ego. Jednocześnie propozycja Golda jest bardzo ogólna — modeluje wyuczalność wskazując na jej ograniczenia, bez odwołania do poszczególnych praktycznych realizacji algorytmów uczących się. Model ten nie wywarł wpływu ani na współczesną psychologię poznawczą ani epistemologię. Stało się tak prawdopodobnie dlatego, że wykorzystuje on zaawansowany matematycznie zestaw narzędzi. Propozycja ta spotkała się z potężnym odzewem w informatyce teoretycznej, językoznawstwie komputerowym i badaniach nad sztuczną inteligencją. Jej aspekt psycholingwistyczny pozostał praktycznie niezbadany.

2. Podstawowe definicje

Poniżej wprowadzam terminologię potrzebną do dalszych rozważań. *Alfabet* to skończony i niepusty zbiór symboli.

Na przykład: $A = \{a, b\}$ — alfabet binarny;

Słowo to skończony ciąg symboli wybranych z pewnego alfabetu.

Na przykład: ciąg „aaabbbbaababaa” jest słowem nad alfabetem A .

Jeśli A jest alfabetem, to przez A^k oznaczamy zbiór słów długości k nad alfabetem A . Na przykład $\{a, b\}^2 = \{aa, ab, ba, bb\}$. Dla dowolnego alfabetu A mamy $A^0 = \{\lambda\}$.

Zbiór wszystkich słów nad alfabetem A oznaczamy przez A^* , np. $\{a, b\}^* = \{\lambda, a, b, aa, ab, ba, bb, aaa, \dots\}$. Innymi słowy $A^* = \bigcup_{n \in \omega} A^n$.

Dowolny zbiór słów wybranych z A^* , dla pewnego alfabetu A , nazywamy *językiem*. Jeśli A jest alfabetem oraz $L \subseteq A^*$, to mówimy, że L jest językiem nad A^* .

Definicja 1. Gramatyka G jest to struktura postaci (A, Σ, S, P) , gdzie:

- A jest alfabetem (alfabetem terminalnym);
- Σ jest zbiorem zmiennych (alfabetem nieterminalnym);
- $S \in \Sigma$ jest symbolem początkowym;
- P jest skończonym zbiorem par postaci $\alpha_i \longrightarrow \beta_i$ dla $i = 1, \dots, n$ oraz $\alpha_i, \beta_i \in (A \cup \Sigma)^*$.

Definicja 2. Dla dowolnych $\gamma, \gamma' \in (A \cup \Sigma)^*$ powiemy, że γ' jest w gramatyce G bezpośrednio wyprowadzalna z γ , czyli $\gamma \rightarrow_G \gamma'$ wtedy i tylko wtedy, gdy istnieją η_1, η_2 oraz $i = 1, \dots, n$, takie, że $\gamma = \eta_1 \alpha_i \eta_2$ oraz $\gamma' = \eta_1 \beta_i \eta_2$.

Definicja 3. γ' jest w gramatyce G wyprowadzalna z γ , czyli $\gamma \rightarrow_G^* \gamma'$ wtedy i tylko wtedy, gdy istnieje ciąg $\gamma_1, \dots, \gamma_m \in (A \cup \Sigma)^*$ taki, że $\gamma = \gamma_1$, $\gamma' = \gamma_m$ oraz $\gamma_i \rightarrow_G \gamma_{i+1}$, dla $i = 1, \dots, m-1$.

Definicja 4. Język generowany przez gramatykę G jest to zbiór: $L(G) = \{\gamma \in A^* : S \rightarrow_G^* \gamma\}$.

Hierarchia języków

- Język L jest rekurencyjnie przeliczalny wtedy i tylko wtedy, gdy istnieje maszyna Turinga e taka, że

$$\alpha \in L \iff \{e\}(\alpha) \downarrow, \text{ dla dowolnego } \alpha \in A^*,$$

czyli: maszyna Turinga e zatrzymuje się tylko na słowach należących do języka L .

- Język L jest rekurencyjny wtedy i tylko wtedy, gdy istnieje maszyna Turinga M licząca funkcję χ_L taką, że dla dowolnego $\alpha \in L$:

$$\chi_L(\alpha) = \begin{cases} 0 & \text{jeśli } \alpha \notin L \\ 1 & \text{jeśli } \alpha \in L. \end{cases}$$

- Język L jest pierwotnie rekurencyjny wtedy i tylko wtedy, gdy jego funkcja charakterystyczna jest pierwotnie rekurencyjna (w sprawie pojęć teorii rekursji zob. (Shoenfield 1991), (Cutland 1980)).
- Języki kontekstowe są postaci $L(G)$, gdzie wszystkie produkcje gramatyki G są postaci:

$$\eta_1 Y \eta_2 \rightarrow_G \eta_1 \beta \eta_2, \text{ dla } Y \in \Sigma, \eta_1, \eta_2, \beta \in (A \cup \Sigma)^*.$$

- Języki bezkontekstowe są postaci $L(G)$, gdzie wszystkie produkcje gramatyki G są postaci:

$$Y \rightarrow_G \beta, \text{ dla } Y \in \Sigma, \beta \in (A \cup \Sigma)^*.$$

- Języki regularne są postaci $L(G)$, gdzie wszystkie produkcje gramatyki G są postaci:

$$Y \rightarrow_G \alpha Z \text{ lub } Y \rightarrow_G \alpha, \text{ dla } Y, Z \in \Sigma, \alpha \in A^*.$$

3. Identyfikowalność w granicy

Jak wcześniej zaznaczyłam, człowiek nie uczy się pierwszego języka poprzez zapamiętywanie reguł gramatycznych podanych *explicite*. Odbywa się to przez pewien rodzaj generalizacji zasłyszanych elementów języka. Od algorytmu uczącego się oczekujemy tego samego — przyswajania zbioru reguł gramatycznych na podstawie szczątkowych danych, przypadkowych elementów języka.

Model identyfikowalności w granicy (Gold 1967) ukazuje możliwy przebieg tego procesu. Model działa dla ustalonej klasy języków. Zostaje z niej wybrany jeden język. Uczeń otrzymuje informacje na jego temat. Ukazane są one na jeden z możliwych sposobów — według przyjętej metody prezentacji danych. Zadanie ucznia polega na odgadnięciu nazwy wybranego języka. Nazwami języków są ich opisy, na przykład gramatyki. Model identyfikacji w granicy polega więc na wyszukiwaniu odpowiedniej gramatyki opisującej zaprezentowany ciąg informacji dotyczącej wybranego języka. W każdym kolejnym kroku procedury uczniowi prezentowana jest jednostka danych dotycząca nieznanego języka. Każdorazowo uczeń ma więc do dyspozycji tylko skończony zbiór danych. W każdym kroku uczeń wybiera nazwę języka. Ta procedura przebiega w nieskończoność. Język jest identyfikowany w granicy, gdy po pewnym czasie uczeń wybiera poprawnie ciągle tę samą nazwę (gramatykę). Właściwa identyfikowalność dotyczy klas języków. Cała klasa języków jest identyfikowalna w granicy, gdy istnieje algorytm zgadujący (uczeń) taki, że identyfikuje on w granicy dowolny język z tej klasy. Warto zaznaczyć, że uczeń nie ma dostępu do informacji, kiedy jego zgadnięcia są prawidłowe. Przetwarza informacje w nieskończoność, ponieważ nie może stwierdzić, czy w kolejnym kroku nie dostanie informacji, która zmusi go do zmiany wyboru.

Identyfikowalność zależy trzech czynników: wybranej klasy języków, metody prezentacji danych dotyczących języka oraz przyjętej relacji nazywania. Najważniejszą częścią modelu jest zaś algorytm zgadujący (uczeń).

Klasa języków Zanim uruchomimy procedurę uczenia się należy wybrać pewną klasę języków $\Omega = \{L : L \subseteq A^*\}$. Będziemy w pracy rozważać klasy języków skończonych, regularnych, bezkontekstowych, kontekstowych, pierwotnie rekurencyjnych, rekurencyjnych i rekurencyjnie przeliczalnych. Klasy te zostały zdefiniowane wcześniej. Z ustalonej klasy wybieramy jeden język, dla ustalenia uwagi L_{44} . Zadaniem ucznia będzie odszukanie nazwy opisującej L_{44} . Wybierać będzie spośród udostępnionego mu zbioru nazw wszystkich języków z klasy Ω .

Metoda prezentacji danych Informacja o języku może być prezentowana uczniowi na dwa sposoby: jako tekst lub jako informant. Zdefiniowane poniżej ciągi danych nazywamy inaczej ciągami treningowymi.

Definicja 5. TEKST dla L jest ω – ciągiem I słów $\alpha_1, \alpha_2, \dots \in L$, takim, że każde słowo $\alpha \in L$ pojawia się przynajmniej raz w I .

Tekst dla danego języka L składa się z elementów L . Dla każdego języka jest ciągiem nieskończonym, ponieważ dopuszczamy powtarzanie się jego elementów. Identyfikowalność przy pomocy tekstu nazywana jest też identyfikowalnością z danych pozytywnych.

W zależności od sposobu wyliczenia i -tego elementu tekstu wyróżniamy następujące rodzaje tekstu:

- arbitralny, to taki, którego i -ty element wyliczany jest przez dowolną funkcję;
- rekurencyjny, to taki, którego i -ty element wyliczany jest przez dowolną rekurencyjną funkcję od numeru kroku w którym element się pojawia;
- pierwotnie rekurencyjny, to taki, którego i -ty element jest wyliczany przez dowolną pierwotnie rekurencyjną funkcję od numeru kroku w którym element się pojawia.

Definicja 6. INFORMANT dla L jest ω – ciągiem I elementów z $(A^* \times \{0, 1\})$, takim, że dla każdego $\alpha \in A^*$:

$$\begin{aligned} (\alpha, 1) &\in I && \text{jeśli } \alpha \in L \\ (\alpha, 0) &\in I && \text{jeśli } \alpha \notin L. \end{aligned}$$

Informant to uporządkowanie wszystkich możliwych ciągów nad danym alfabetem wraz z informacjami, czy elementy te należą, czy nie należą do danego języka. Tak więc, I jest identyczna z funkcją charakterystyczną χ_L zbioru L względem zbioru A^* . Identyfikowalność przy pomocy informantu nazywana jest też identyfikowalnością poprzez dane pozytywne i negatywne. W zależności od sposobu wyliczenia i -tego ciągu wyróżniamy następujące rodzaje informantu:

- arbitralny, to taki, którego i -ty element jest wyliczany przez dowolną funkcję, o ile każdy ciąg z A^* pojawia się w informancie przynajmniej raz;
- metodyczny, to taki, który jest apriorycznym uporządkowaniem A^* , zaś i -ty element to element w porządku o numerze i ;
- żądający, to taki, którego i -ty element wybierany jest przez ucznia w kroku i na podstawie dotychczasowych danych.

Relacja nazywania Uczeń może używać dwóch tzw. relacji nazywania, czyli dwóch sposobów decydowania o adekwatności danej nazwy dla poszukiwanego języka:

Definicja 7. Generator dla L jest Maszyną Turinga e taką, że:

$$L = \{\alpha \in A^* : \{e\}(\alpha) \downarrow\}.$$

Jak już wspomnieliśmy nazwy mają charakter opisowy. Z gramatyki – nazwy opisowej danego języka, łatwo otrzymujemy maszynę Turinga będącą generatorem. Nazywanie przebiega poprzez generowanie z wybranej gramatyki odpowiednich ciągów oraz porównywanie ich z otrzymanymi wcześniej danymi. Jeśli są zgodne, uczeń decyduje się na wybór tej maszyny Turinga. Jeśli nie, wybiera inną i z nią postępuje tak samo.

Definicja 8. Tester dla L jest Maszyną Turinga, która oblicza funkcję χ_L taką, że dla każdego $\alpha \in A^*$:

$$\chi_L(\alpha) = \begin{cases} 0 & \text{jeśli } \alpha \notin L \\ 1 & \text{jeśli } \alpha \in L. \end{cases}$$

Nazywanie przebiega w tym przypadku poprzez sprawdzanie wybranym testerem przynależności do języka wszystkich elementów otrzymanego ciągu danych. Jeśli dla wszystkich elementów odpowiedź testera była pozytywna, uczeń wybiera ten tester. Jeśli dla któregoś wynik postępowania był negatywny, uczeń odrzuca tę hipotezę i wybiera inny tester.

Algorytm zgadujący Algorytmowi zgadującemu odpowiada funkcja $Zg : I \longrightarrow \mathcal{N}$, gdzie I jest zbiorem wszystkich możliwych skończonych ciągów informacji, $\mathcal{N}$ jest zbiorem wszystkich nazw języków z danej klasy.

Definicja 9. Niech $I_L = i_1, i_2, i_3, \dots$ będzie ω -ciągiem treningowym dla L . Wtedy:

L jest identyfikowany w granicy przez funkcję Zg względem ciągu $I_L \iff$ istnieje nazwa N języka L taka, że $\exists k \forall u \geq k \Rightarrow Zg(i_1, \dots, i_u) = N$.

Definicja 10. Niech Ω będzie klasą par postaci (L, N) , gdzie N jest nazwą języka L . Wtedy:

Ω jest identyfikowalna w granicy $\iff$ dowolny język L z klasy Ω jest identyfikowany w granicy (oczywiście względem odpowiedniego systemu nazw).

Z reguły ważniejsze klasy języków mają pewne kanoniczne systemy nazw, na przykład: języki regularne można nazywać za pomocą automatów skończonych, gramatyk regularnych lub wyrażeń regularnych. Odpowiednie systemy nazw są tutaj wzajemnie efektywnie przekładalne. Nie ma więc znaczenia,

który system wybieramy. Podobne systemy nazw mają języki bezkontekstowe (gramatyki bezkontekstowe lub automaty ze stosem), języki pierwotnie rekurencyjne (terminy definiujące pierwotnie rekurencyjne funkcje charakterystyczne). Maszyny Turinga liczące funkcje charakterystyczne są kanonicznym systemem nazw dla języków rekurencyjnych. W tym przypadku sam system nazw nie jest jednak rekurencyjny.

W pracach poświęconych teorii uczenia przyjmuje się, że każda z tych podstawowych klas języków zadana jest wraz z odpowiadającym jej kanonicznym systemem nazw. Klasy języków rozpatrywane z punktu widzenia wyuczalności muszą być reprezentowane wraz z odpowiednim systemem nazw, czyli przyporządkowaniem każdemu językowi zbioru jego nazw lub, równoważnie, należy rozpatrywać klasy par (L, N) , gdzie N jest nazwą języka L . Przypomnijmy, że nazwa zawiera skończony opis odpowiedniego języka. Przyjęte jest, że tam, gdzie istnieje taki kanoniczny system nazw pomija się go w opisie odpowiedniej klasy języków. Tak więc, mówimy po prostu o wyuczalności klas języków regularnych, bezkontekstowych, kontekstowych, pierwotnie rekurencyjnych, rekurencyjnych, milcząco zakładając, że określony jest odpowiedni kanoniczny system nazw.

Przyjmując powyższy aparat możemy udowodnić następujące twierdzenia (Gold 1967):

Twierdzenie 1. *Następujące warunki są równoważne:*

1. *klasa języków jest identyfikowalna w granicy przy pomocy arbitralnego informantu;*
2. *klasa języków jest identyfikowalna w granicy przy pomocy metodycznego informantu;*
3. *klasa języków jest identyfikowalna w granicy przy pomocy żądającego informantu.*

Twierdzenie 2. *Klasa języków rekurencyjnych nie jest identyfikowalna w granicy przez metodyczny informant z generatorem.*

Twierdzenie 3. *Klasa języków skończonych jest identyfikowalna w granicy przez arbitralny tekst z testerem.*

Twierdzenie 4. *Żadna klasa języków zawierająca wszystkie języki skończone i przynajmniej jeden nieskończony (tzw. nadskończona klasa języków) nie jest identyfikowalna w granicy przez rekurencyjny tekst z generatorem.*

Poniżej znajduje się zestawienie dotyczące wyuczalności w modelu identyfikowalności w granicy.

Pierwotnie rekurencyjny tekst z generatorem	rekurencyjnie przeliczalne rekurencyjne
Informant	pierwotnie rekurencyjne kontekstowe bezkontekstowe regularne nadskończona
Tekst	języki skończonej mocy

Tabela 1. Wyuczalność i niewyuczalność klas języków w modelu Golda

4. Psycholingwistyczna analiza modelu Golda

Przejdźmy teraz do wyodrębnienia i analizy podstawowych założeń modelu Golda. Porównamy je z naturalnym procesem przyswajania zdolności językowych.

Uczenie jest procesem nieskończonym Założenie to odpowiada intuicji na temat nabywania zdolności językowych. Nie tylko w dzieciństwie rozwijamy zdolności lingwistyczne. Edukacja językowa nie kończy się w jakimś ustalonym momencie, że wciąż dowiadujemy się nowych rzeczy na temat języka. Z czasem dane językowe coraz rzadziej są istotne dla rozwoju naszej kompetencji językowej. Nawet jeśli uznamy, że proces nabywania wiedzy składniowej jest skończony i upływa wraz z np. trzynastym rokiem życia (zob. (Kurcz 2000)), to wciąż nieskończony model identyfikacji możemy z powodzeniem odnosić do procesu przyswajania semantycznej wiedzy językowej. Ponadto tak jak w naturalnym procesie tak i w modelu Golda uczeń nigdy nie wie, kiedy gramatyka, którą się posługuje jest „prawidłowa”. Wciąż istnieje bowiem możliwość, że niektóre z danych językowych spowodują zmiany w przyswojonej gramatyce.

Efektem uczenia się jest ustabilizowanie ciągu wyników na jednej, poprawnej odpowiedzi Założenie to można osłabić oczekując jedynie stabilizacji odpowiedzi na jednym języku, dopuszczając wahania w identyfikacji konkretnej gramatyki. Osłabienie definicji identyfikowalności w granicy może zmienić własności wyuczalności. Co więcej, niewykluczone, że może doprowadzić do nowych wyników w dziedzinie teorii uczenia się.

Ponadto, uznajemy czyjaś znajomość języka na podstawie sprawnego porozumiewania się. Na drugi plan schodzi wtedy kwestia, jaką gramatyką dany osobnik się posługuje. Nie musimy zakładać, że język naturalny posiada jeden prawidłowy opis gramatyczny i że właśnie ten opis jest standardowo

przyswajany przez wszystkich użytkowników języka. Co więcej, nie musimy nawet zakładać, że pojedynczy człowiek ustala swą wiedzę językową raz na zawsze. Stabilizacja ta nie jest konieczna, o ile jej brak nie zakłóca sprawnego porozumiewania się. Takie złagodzenie warunku przyswojenia języka nie jest zgodne ze stanowiskiem Chomsky'ego. Według niego język posiada jedną prawidłową strukturę gramatyczną, która każdorazowo podlega wyuczeniu (Chomsky 1995) (por. (Quine 1990), (Gierasimczuk 2004)).

Pojęcie identyfikowalności dotyczy określonych klas języków Jeżeli uczenie miałooby przebiegać na zasadzie identyfikacji języka wybranego z pewnej określonej klasy, mielibyśmy poszlakę w postaci górnej granicy złożoności języka naturalnego. Na przykład, jeśli mechanizm uczenia umożliwia wyuczalność tylko klasy języków kontekstowych, to górnym ograniczeniem na złożoność struktury języka naturalnego byłaby właśnie kontekstowość. W drugą stronę ta zależność nie zachodzi. Przyjmijmy, że udało się przeanalizować dostępne języki naturalne pod kątem ich struktury składniowej, a następnie wydać sąd na temat górnego ograniczenia ich złożoności. Wciąż możliwe jest to, że mechanizm wyuczalności działa również dla języków z wyższych klas. Co więcej możliwe, że ten sam mechanizm wyuczalności stosuje się do rozmaitych umiejętności pozajęzykowych. Uczymy się przecież nie tylko języka, lecz także gestykulacji, mimiki czy prowadzenia samochodu. Aktywności tego rodzaju mogą mieć dowolną złożoność, której nie jesteśmy w stanie w tej chwili oszacować. Daleko nam bowiem do zrozumienia i matematycznego opisanie zasad rządzących tymi zdolnościami. Dlatego możliwe, że mechanizm wyuczalności jest silniejszy niż oczekivalibyśmy, gdyby dotyczył tylko uczenia się języka naturalnego.

Uczeń posiada umiejętność testowania ciągów za pomocą automatów i generowania ciągów za pomocą gramatyk Jest to założenie dotyczące struktury algorytmu zgadującego (relacji nazywania), które można rozważyć w kontekście wrodzonych mechanizmów poznawczych. Nie wydaje się ono zbyt kontrowersyjne. Można je bez większych zastrzeżeń akceptować.

Zaletą propozycji Golda jest fakt, że model ten nie wymaga istnienia jakichkolwiek wrodzonych uczniowi struktur gramatycznych. Struktura gramatyczna jest w całości przyswajana z próbek językowych przez pewnego rodzaju uogólnienie. Przyjmując model Golda pomija się więc kwestię wspólnego zrębu wszystkich języków ludzkich — tzw. gramatyki uniwersalnej. W założeniu miałyby ona być tym, co łączy wszystkie języki ludzkie, a zarazem wszystkie ludzkie umysły w aspekcie kompetencji językowej (Chomsky 1995).

Ciekawym wynikiem pracy *Identification in the Limit* jest twierdzenie 4. Głosi ono, że żadna klasa, która mogłaby zawierać język naturalny, nie jest wyuczalna z danych pozytywnych (tekstu). Język naturalny zwykle umieszcza się w hierarchii na wysokości języków bezkontekstowych, niekiedy kontekstowych (zob. (Partee et al. 1993)). Przyjęcie modelu Golda jako trafnie opisującego mechanizm nabywania języka daje niebanalny wynik istotny dla psychologii poznawczej, epistemologii, językoznawstwa, jak również badań nad sztuczną inteligencją. Okazuje się bowiem, że dla identyfikacji (wyuczalności) języka naturalnego potrzeba czegoś więcej niż tylko danych pozytywnych. Wymaga ona informacji na temat granic wyuczanego języka w postaci informacji negatywnej: *to zdanie nie należy do poszukiwanego języka*. Czy wynik ten zgodny jest ze świadectwami empirycznymi dotyczącymi naturalnego przyswajania języka? Niektórzy uważają, że informacja negatywna jest znikoma i ignorowana przez dziecko (Chomsky 1965). Badacze wciąż spierają się na ten temat, proponując nowe ograniczenia na model wyuczalności tak, aby rozszerzyć identyfikowalność z danych pozytywnych na klasy języków wyższe niż skończone (zob. np. (Angluin 1980)). Ograniczenie wynikające z twierdzenia 4 nie wydaje nam się kontrowersyjnie. Dziecko we wczesnym stadium lingwistycznym otrzymuje dane pozytywne w postaci zasłyszanych zdań, negatywne zaś w postaci korekt pochodzących od nauczyciela. Pojawia się tutaj pytanie o nieodzowność udziału nauczyciela w procesie przyswajania języka.

5. Uczenie się semantyki

Dotychczas zastanawialiśmy się jedynie nad wyuczalnością składniowego aspektu języka. Naturalnym kierunkiem dalszych rozważań jest kwestia uczenia się semantyki języka¹. Problem ten jest stosunkowo słabo zbadany. Jako przykład konstrukcji semantycznych badanych pod względem wyuczalności podać można konstrukcje kwantyfikatorowe. Problem wyuczalności konstrukcji kwantyfikatorowych po raz pierwszy podjęty został przez Johana van Benthema (van Benthem 1986), a następnie rozważany w kilku innych pracach (zob. (Clark 1996), (Florêncio 2002), (Tiede 1999)). Były to próby zastosowania modelu uczenia się składni zaprezentowanego w niniejszej pracy do problemu uczenia się semantyki języka.

W jaki sposób przyswajamy znaczenie wyrażeń kwantyfikatorowych w języku naturalnym? Znaczenie wyrażeń języka naturalnego można utożsamiać z procedurami rozpoznawania ich ekstensji (Moschovakis 1994), (Szymanik

¹ Zagadnienie wyuczalności semantyki obliczeniowej potraktowane zostało tutaj skrótowo. Więcej szczegółów znajdzie czytelnik w (Gierasimczuk 2007).

2004). Rozsądnym wydaje się przyjęcie, że przyswajanie znaczenia wyrażeń kwantyfikatorowych odbywa się poprzez uczenie się procedur wyznaczania ich denotacji. Opierając się na definicji kwantyfikatora monadycznego zaproponowanej przez Lindströma (Lindström 1966), możemy zakodować klasy modeli skończonych odpowiadających danemu kwantyfikatorowi. Takie klasy modeli skończonych mogą być reprezentowane przez odpowiednie języki (Mostowski 1998). Jako takie, mogą podlegać wyuczalności w tym samym sensie, co języki w ujęciu składniowym.

Poniżej znajduje się kilka faktów dotyczących klas kwantyfikatorów monadycznych i modeli obliczeń rozpoznających te klasy.

- Kwantyfikator Q jest definiowalny w logice pierwszego rzędu $\iff L_Q$ jest akceptowany przez pewien acykliczny automat skończony (van Benthem 1986).
- Monadyczny kwantyfikator Q jest definiowalny w logice podzielności $FO(D_\omega) \iff L_Q$ jest akceptowany przez pewien automat skończony (Mostowski 1998).
- W języku naturalnym istnieje wiele kwantyfikatorów niedefiniowalnych środkami gramatyk bezkontekstowych.

Oto kilka rezultatów połączenia idei kodowania konstrukcji kwantyfikatorowych oraz wyuczalności składniowej.

- Klasy FO , $FO(D_\omega)$ oraz pólliniowych kwantyfikatorów nie są identyfikowalne w granicy z danymi pozytywnymi, ale są identyfikowalne w granicy z danymi pozytywnymi i negatywnymi.
- Istnieją podklasy klas wymienionych powyżej, które są identyfikowalne w granicy przy pomocy tekstu, np. zbiór kwantyfikatorów *left upward monotone* (Tiede 1999).

Kwestia adekwatności narzędzi teorii uczenia się składni do opisu uczenia się semantyki jest oczywiście dyskusyjna. Wątpliwości mogą dotyczyć adekwatności tego aparatu do kształtowania kompetencji semantycznej. W obrębie tej kompetencji wyróżnić można kilka powiązanych ze sobą mechanizmów: zdolność do sprawdzania wartości logicznej zdania w danym modelu, zdolność do rozpoznawania relacji wynikania pewnego zdania z innych zdań oraz zdolność do generowania adekwatnych opisów, czyli zdań prawdziwych w danym modelu (Mostowski 1994), (Bucholc 2004), (Mostowski Wojtyniak 2004). Można więc kompetencję semantyczną przedstawić w następujący sposób.

Faktem jest, że rozpatrywane wcześniej modele wyuczalności nie są wrażliwe na rozróżnienie pomiędzy rozpoznawaniem a generowaniem adekwatnych opisów. Można argumentować, że istnieje wzajemny przekład pomiędzy automatami a gramatykami (zob. (Hopcroft et al. 2001)). Nie wydaje

Kompetencja semantyczna	zdolność sprawdzania wartości logicznej
	wejście: model M oraz zdanie φ ; Czy $M \models \varphi$?
	zdolność do rozpoznawania
	relacji wynikania
	wejście: zdania φ, ψ ; Czy $\varphi \models \psi$?
	zdolność do generowania
	adekwatnych opisów
	wejście: model M ; znajdź zdanie φ takie, że $M \models \varphi$

Rysunek 1. Kompetencja semantyczna

się jednak aby taka redukcja poprawnie zdawała sprawę z procesu przyswajania konstrukcji semantycznych. Wydaje się, że najpierw człowiek rozumie konstrukcje semantyczne, a dopiero później sam zaczyna je konstruować. Generowanie opisów jest prawdopodobnie bardziej skomplikowane niż testowanie ich trafności (zob. (Bates 1995), (Benedict 1979), (Clark 1993), (Fraser 1979), (Goldin-Meadow 1976), (Layton 1979)).

6. Podsumowanie

W konfrontacji z naturalnym procesem przyswajania języka model identyfikacji w granicy szczególnie dobrze oddaje to, że uczeń nigdy nie dowiadyje się kiedy przyswojona przez niego gramatyka jest „prawidłowa”. Inną zaletą metodologiczną modelu Golda jest fakt, że pojęcie identyfikowalności dotyczy określonych klas języków. Założenie to pozwala wnioskować o strukturze języka naturalnego (jego miejscu w hierarchii języków) na podstawie struktury kompetencji językowej.

Ważną cechą zaprezentowanego modelu jest to, że nie zakłada on istnienia wrodzonego zrębu wszystkich gramatyk, w sensie gramatyki uniwersalnej postulowanej przez Chomsky’ego. Okazuje się więc, że przyjęcie algorytmicznego modelu kompetencji językowej nie wymaga postulowania części wspólnej wszystkich języków naturalnych. Ich podobieństwo może być spowodowane przynależnością do takiej klasy, którą wrodzony mechanizm uczenia się jest w stanie identyfikować, na przykład klasy wszystkich języków bezkontekstowych.

Zagadnieniem ściśle wiążącym się z wyuczalnością składni języka jest wyuczalność semantyki. Algorytmiczna teoria znaczenia stwarza możliwość zredukowania problemu przyswajania znaczeń konstrukcji semantycznych do

problemu przyswajania odpowiadających tym konstrukcjom języków. Języki te uzyskiwane są poprzez kodowanie modeli skończonych, w których zdanie z odpowiednią konstrukcją jest prawdziwe. Dzięki takiemu zabiegowi możemy stosować narzędzia formalnej teorii uczenia się, powstałej na potrzeby badań składniowych, do problemu uczenia się semantyki. Przykładem takiego narzędzia może być *model identyfikacji w granicy*. Pozostaje jednak pytanie o adekwatność takiego podejścia wobec złożoności naturalnego procesu przyswajania znaczeń z jednej strony, a z drugiej — złożoności postulowanego modelu kompetencji semantycznej.

Literatura

- Angluin, D. (1980). Inductive inference of formal languages from positive data. *Information and Control* **45**.
- van Benthem, J. (1986). *Essays in logical semantics*. Dordrecht: Reidel Publishing Company.
- Bates, E., P.S. Dale and D. Thal (1995). *Individual Differences and their Implications*, W: P. Fletcher, B. MacWhinney (red.), *The Handbook of Child Language*. Oxford: Blackwell.
- Benedict, H. (1979). Early Lexical Development: Comprehension and Production. *Journal of Child Language* **6**.
- Bucholc, P. (2004). *Kompetencja logiczna a poprawność logiczna. Analiza na przykładzie terminów pustych*, W: J. Szymanik, M. Zajenkowski (red.), *Kognitywistyka. O umyśle umyślnie i nieumyślnie*. Warszawa: Koło Filozoficzne przy Kolegium MISH.
- Chomsky, N. (1959). Review of Verbal Behavior by B.F. Skinner. *Language* **35/1**.
- (1965). *Aspects of the theory of syntax*. Cambridge: The M.I.T. Press.
- (1965). *Cartesian linguistics*. Nowy Jork: Harper & Row.
- (1968). *Language and mind*. Nowy Jork: Harcourt Brace.
- (1967). Recent contributions to the theory of innate ideas. *In Syntheses* **17**.
- (1995). *The minimalist program*. Cambridge: The MIT Press.
- Clark, E.V. (1993). *The Lexicon in Acquisition*. Cambridge: Cambridge University Press.
- Clark, R. (1996). Learning first-order quantifiers denotations. An essay in semantic learnability. *IRCS Technical Report, University of Pennsylvania*.
- Cutland, N. (1980). *Computability. An introduction to recursive function theory*. Cambridge: Cambridge University Press.
- Fraser, C., U. Bellugi (1979). Control of Grammar in Imitation, Production, and Comprehension. *Journal of Verbal Learning and Verbal Behaviour* **2**.
- Florêncio, C. C. (2002). Learning generalized quantifiers. *Proceedings of Seventh ESSLLI Student Session*.
- Gierasimczuk, N. (2004). *Teoretyczny model nabywania języka według Quine'a*, W: J. Szymanik, M. Zajenkowski (red.), *Kognitywistyka. O umyśle umyślnie i nieumyślnie*. Warszawa: Koło Filozoficzne przy Kolegium MISH.

- (2007). *The Problem of Learning the Semantics of Quantifiers*, W: B. ten Cate, H. Zeevat (red.), Proceedings of the Tbilisi Symposium on Language, Logic and Computation 2005, Lectures Notes in Artificial Intelligence 4363.
- Gold, E. M. (1967). Language identification in the limit. *Information and Control* **10**.
- Goldin-Meadow, S., M.E.P. Seligman (1976). Language in the Two Year Old. *Cognition* **4**.
- Hopcroft, J., R. Motwani, J. Ullman (2001). *Introduction to automata theory, languages, and computation*. Nowy Jork: Addison-Wesley.
- Kurcz, I. (2000). *Psychologia języka i komunikacji*. Warszawa: Wydawnictwo Naukowe „Scholar”.
- Layton, T.L., S.L. Stick (1979). Comprehension and Production of Comparatives and Superlatives. *Journal of Child Language* **6**.
- Lindström, P. (1966). First order predicate logic with generalized quantifiers. *Theoria* **32**.
- Lyons, J. (1977). *Semantics*. Cambridge: Cambridge University Press.
- Macnamara, J. (1986). *A Border dispute. The place of logic in psychology*. Cambridge: The MIT Press.
- Moschovakis (1994). *Sense and denotation as algorithm and value*, W: J. Oikkonen, J. Väänänen (red.), Lecture Notes in Logic 2. Berlin: Springer-Verlag.
- Mostowski, M. (1994). *Kwantyfikatory rozgałęzione a problem formy logicznej*, W: M. Omyła (red.), Nauka i język. Warszawa: Biblioteka Myśli Semiotycznej.
- (1998). Computational semantics for monadic quantifiers. *Journal of Applied Non-Classical Logics* **8**.
- Mostowski, M., D. Wojtyński (2004). Computational Complexity of the Semantics of Some Natural Language Constructions. *Annals of Pure and Applied Logic* **127**.
- Partee, B., A. ter Meulen, R. E. Wall (1993). *Mathematical methods in linguistics*. Dordrecht: Kluwer Academic Publishers.
- Pinker, S. (2000). *The language instinct. The new science of language and mind*. London: Penguin Books.
- Quine, W. V. O. (1990). *Pursuit of Truth*. Cambridge: Harvard University Press.
- Shoenfield, J. R. (1991). *Recursion theory*. Berlin: Springer-Verlag.
- Szymanik, J. (2004). Semantyka obliczeniowa dla kwantyfikatorów monadycznych w języku naturalnym, Praca magisterska w Instytucie Filozofii UW, promotor: prof. Marcin Mostowski.
- Tiede, H. J. (1999). Identifiability in the limit of context-free generalized quantifiers. *Journal of Language and Computation* **1**.